

Evaluating Fairness and Algorithmic Bias in Ensemble-Based Diabetes Risk Prediction Models Across Diverse Socio-Demographic and Ethno-Racial Subpopulations

Authors

Abisola Bashiru, Monique Amama, Grace Ekene, Okeke Gift, Asher Noah

Date; June 28, 2026

Abstract

The integration of artificial intelligence into clinical risk prediction has demonstrated significant promise for early diabetes detection, yet mounting evidence reveals that algorithmic models can inadvertently perpetuate or exacerbate health disparities across socio-demographic and ethno-racial groups. This study addresses the critical gap in fairness evaluation within ensemble-based diabetes risk prediction by systematically assessing algorithmic bias across diverse subpopulations. Utilizing the PIMA Indian Diabetes dataset augmented with synthetic socio-demographic attributes, we implemented and evaluated multiple ensemble learning architectures including AdaBoost, Gradient Boosting, XGBoost, and a Stacking Ensemble, with fairness metrics including disparate impact, statistical parity difference, and equalized odds. The Stacking Ensemble achieved the highest overall predictive performance with 86.0% accuracy and 0.91 ROC-AUC ; however, significant disparities emerged across subpopulations, with false negative

rates disproportionately affecting minority groups by margins of 12-18%. Fairness-aware optimization reduced these disparities by 42% while maintaining 83.4% accuracy. These findings demonstrate that high predictive performance does not guarantee equitable outcomes, underscoring the necessity of integrating fairness constraints throughout the model development lifecycle to ensure that AI-driven diabetes risk assessment serves all populations equitably.

Keywords: Algorithmic Fairness, Ensemble Learning, Diabetes Risk Prediction, Health Equity, Bias Mitigation

1. Introduction

1.1 Background

Diabetes represents one of the most pressing global health challenges of the twenty-first century, affecting an estimated 529 million adults worldwide with projections reaching 1.3 billion by 2050 . The disease disproportionately impacts racial and ethnic minorities and individuals experiencing social disadvantages, with social determinants of health (SDoH)—including education, income, housing stability, and neighborhood deprivation—recognized as key drivers of these persistent disparities . The increasing availability of electronic health records (EHRs) and the rapid advancement of machine learning (ML) techniques have created unprecedented opportunities for developing personalized risk prediction tools. Ensemble learning methods, which combine multiple base learners to improve predictive performance, have demonstrated particular promise in diabetes risk assessment, with stacking ensemble models achieving up to 86% accuracy .

However, the deployment of AI in healthcare has raised substantial concerns regarding algorithmic fairness. Machine learning models trained on observational data that reflects historical inequities can inadvertently perpetuate or amplify existing health disparities . A landmark study revealed that a widely used commercial risk prediction tool exhibited significantly lower accuracy for Black patients compared to white patients, resulting in systematically fewer Black patients being recommended for essential care services . Similarly, fairness evaluations of established clinical prediction models, including the QPrediction family of models (QRISK3, QDiabetes, QStroke), have revealed that prediction parity was not achieved across ethnicity, education level, income, and deprivation indicators . These findings underscore the critical importance of systematically evaluating and mitigating algorithmic bias in clinical prediction models.

1.2 Problem Statement

Despite the growing adoption of ensemble-based machine learning models for diabetes risk prediction, significant gaps exist in the systematic evaluation of algorithmic fairness across diverse socio-demographic and ethno-racial subpopulations. Existing research has largely prioritized predictive accuracy as the primary metric of model performance, with limited attention to how these models perform across different demographic groups . A systematic review of ML-based phenotyping studies found that only 5% assessed fairness, and fairness evaluation remains inconsistently adopted in clinical ML research .

Several critical challenges persist in the current landscape. First, most diabetes prediction studies evaluate overall model performance metrics such as accuracy and AUC without disaggregating results by demographic subgroups, masking potentially significant disparities in predictive performance. Second, when fairness is considered, it is often treated as an afterthought rather than integrated throughout the model development pipeline . Third, the comparative effectiveness of different fairness-aware strategies—including pre-processing techniques (e.g., reweighting, oversampling), in-processing methods (e.g., adversarial debiasing, fairness constraints), and post-processing approaches—remains poorly understood in the context of ensemble-based diabetes prediction . Fourth, the intersectional nature of health disparities, where individuals belonging to multiple marginalized groups face compound disadvantages, is rarely addressed in fairness evaluations .

The central problem this study addresses is the absence of a comprehensive framework for evaluating and mitigating algorithmic bias in ensemble-based diabetes risk prediction models across diverse socio-demographic and ethno-racial subpopulations.

1.3 Objectives of the Study

General objective:

To develop and validate a comprehensive framework for evaluating and mitigating algorithmic bias in ensemble-based diabetes risk prediction models across diverse socio-demographic and ethno-racial subpopulations.

Specific objectives:

1. To implement and systematically evaluate multiple ensemble learning architectures (AdaBoost, Gradient Boosting, XGBoost, and Stacking Ensemble) for diabetes risk prediction.
2. To assess the fairness of these models across diverse demographic subgroups using established fairness metrics including disparate impact, statistical parity difference, equalized odds, and false negative rate ratios.
3. To evaluate the comparative effectiveness of three fairness-aware strategies—pre-processing (oversampling/reweighting), in-processing (adversarial debiasing), and post-

processing (threshold adjustment)—in mitigating algorithmic bias while preserving predictive performance.

4. To identify the socio-demographic and clinical predictors that contribute most significantly to algorithmic disparities and propose actionable recommendations for equitable model deployment.

1.4 Research Questions

Research Question 1: To what extent do ensemble-based diabetes risk prediction models exhibit algorithmic bias across socio-demographic and ethno-racial subpopulations, as measured by disparate impact, statistical parity difference, equalized odds, and false negative rate ratios?

Research Question 2: How does the performance of various ensemble architectures (AdaBoost, Gradient Boosting, XGBoost, Stacking Ensemble) differ across subpopulations, and what trade-offs exist between overall predictive accuracy and subgroup-level fairness?

Research Question 3: Which fairness-aware mitigation strategy—pre-processing (oversampling/reweighting), in-processing (adversarial debiasing), or post-processing (threshold adjustment)—most effectively reduces algorithmic bias while maintaining clinically acceptable predictive performance?

Research Question 4: What clinical and socio-demographic features are the primary drivers of algorithmic disparities, and how can these insights inform equitable model design and deployment?

1.5 Significance of the Study

This research makes several significant contributions to clinical AI and health equity.

For practitioners and healthcare administrators, this study provides an evidence-based framework for selecting and evaluating diabetes risk prediction models with explicit consideration of fairness. The comparative analysis of fairness-aware strategies offers actionable guidance for implementing equitable AI systems in clinical settings, with specific metrics to monitor and thresholds to maintain.

For policymakers, the findings highlight the importance of establishing regulatory standards for fairness evaluation in clinical AI, demonstrating that high overall accuracy does not guarantee equitable outcomes. The study provides empirical evidence to support the inclusion of fairness requirements in AI procurement and deployment guidelines.

For academic literature, this research advances the theoretical understanding of algorithmic fairness in healthcare by systematically evaluating fairness across ensemble architectures and mitigation strategies. The integration of bias source analysis—identifying whether disparities originate from data, model, or deployment—contributes to the growing body of knowledge on fairness in clinical ML .

For future researchers, this study establishes a replicable framework for fairness evaluation that can be extended to other clinical prediction tasks and healthcare settings. The methodology for incorporating synthetic socio-demographic attributes to enable fairness analysis in datasets lacking demographic information provides a template for addressing data limitations.

1.6 Scope and Limitations

This study focuses on ensemble-based diabetes risk prediction using the PIMA Indian Diabetes dataset, a widely used benchmark dataset in diabetes research containing 768 samples from female patients aged 21 years and older . Given that the original dataset lacks comprehensive socio-demographic attributes, this study incorporates synthetically generated socio-demographic variables—including ethnicity, socioeconomic status, education level, and insurance type—to enable fairness evaluation across diverse subpopulations. The synthetic attributes are generated based on distributions observed in real-world datasets and validated against known demographic patterns in diabetes prevalence .

The study is limited to the specific ensemble architectures implemented (AdaBoost, Gradient Boosting, XGBoost, Stacking Ensemble) and does not exhaustively evaluate all possible ensemble methods or single-classifier alternatives. The findings are based on a single dataset and may not fully generalize to other populations, healthcare systems, or geographic regions. The synthetic nature of the socio-demographic attributes, while enabling fairness analysis, may not capture the full complexity of real-world health disparities. Additionally, this study focuses on group-level fairness metrics and does not address individual fairness considerations, which represent an important direction for future research .

2. Literature Review

2.1 Conceptual Review

Algorithmic Bias in Healthcare:

Algorithmic bias refers to systematic and unfair discrimination in the outputs of machine learning models, where certain groups are systematically disadvantaged relative to others. In healthcare, algorithmic bias can manifest through multiple mechanisms: biased data (reflecting historical inequities in healthcare access and documentation), biased model design (inappropriate choice of labels or optimization objectives), and biased deployment (models inaccessible or inappropriate for certain populations) . Understanding these sources of bias is essential for developing effective mitigation strategies.

Ensemble Learning:

Ensemble learning is a machine learning paradigm that combines multiple base learners to achieve superior predictive performance compared to individual models. Common ensemble methods include boosting (AdaBoost, Gradient Boosting, XGBoost), which sequentially trains models to correct previous errors, and stacking, which trains a meta-model to combine predictions from diverse base models. Ensemble methods are particularly well-suited for clinical prediction tasks due to their robustness and ability to capture complex patterns in high-dimensional data.

Fairness Metrics in Machine Learning:

Fairness in machine learning is operationalized through various metrics, each capturing different notions of what constitutes fair treatment. **Group fairness** metrics, which this study employs, require that models perform similarly across demographic groups. Key group fairness metrics include:

- **Disparate Impact:** The ratio of the probability of a positive outcome for the protected group to that for the privileged group. A value below 0.8 is generally considered indicative of disparate impact.
- **Statistical Parity Difference:** The difference between the probability of a positive outcome for the protected group and the privileged group. A value close to zero indicates parity.
- **Equalized Odds:** Requires that true positive rates and false positive rates are equal across groups. This ensures that the model's predictive performance is consistent across demographic groups.
- **False Negative Rate Ratio:** The ratio of false negative rates between groups. Disparities in false negative rates are particularly concerning in healthcare, as they may indicate that certain groups are systematically underdiagnosed.

2.2 Theoretical Framework

This study is guided by two complementary theoretical frameworks that together provide a comprehensive lens for understanding and addressing algorithmic fairness in healthcare AI.

Prospect Theory:

Developed by Kahneman and Tversky, Prospect Theory describes how individuals make decisions under uncertainty, with systematic deviations from rational choice. In the context of healthcare AI, Prospect Theory suggests that clinicians may exhibit systematic biases in how they interpret and act upon model predictions, potentially exacerbating algorithmic disparities. For example, clinicians may overweight the risk of false positives for certain patient groups based on stereotypes, or may exhibit "automation bias" by overly trusting model outputs even

when they conflict with clinical judgment . Understanding these cognitive biases is essential for designing fairness-aware AI systems and appropriate clinician training.

The Fairness-Aware Machine Learning Framework:

This framework categorizes fairness interventions into three stages of the ML pipeline :

1. **Pre-processing:** Interventions applied to the data before model training, including reweighting, oversampling, and data augmentation to balance representation across groups.
2. **In-processing:** Interventions integrated into the model training process, including adversarial debiasing, fairness constraints in optimization, and regularization techniques that penalize unfair predictions.
3. **Post-processing:** Interventions applied to model outputs after training, including threshold adjustment and calibration to ensure consistent performance across groups.

This tripartite framework provides a systematic approach to evaluating and implementing fairness-aware machine learning in clinical applications.

2.3 Empirical Review

Bhatta (2026) conducted a comprehensive evaluation of ensemble learning algorithms for diabetes prediction using the PIMA Indian Diabetes dataset. The study compared AdaBoost, Gradient Boosting, XGBoost, and Stacking Ensemble models, finding that the Stacking Ensemble achieved the highest performance with 86.0% accuracy, 0.82 precision, 0.80 recall, 0.81 F1-score, and 0.91 ROC-AUC. While the study provided valuable insights into ensemble model performance, it did not address algorithmic fairness, representing a significant gap that the present study addresses .

Mohd Hamdan et al. (2025) evaluated three fairness-aware approaches—Fairness-Aware Interpretable Modelling (FAIM), Fairness-Aware Machine Learning (FAML), and Fairness-Aware Oversampling (FAWOS)—for diabetes prediction. The study found that FAML achieved perfect fairness metrics while maintaining acceptable accuracy, while FAWOS improved overall classification accuracy without explicitly enforcing fairness. Importantly, the study found that no single method achieved an optimal balance of accuracy, fairness, and interpretability, supporting the need for hybrid approaches. However, the study compared these methods in isolation rather than evaluating their integration into ensemble models .

Huynh et al. (2025) evaluated the fairness of QPrediction risk models (QRISK3, QDiabetes, QStroke) using UK Biobank data, revealing systematic overprediction of risks across all models. The study found that prediction parity was not achieved in false negative rates and true negative rates for ethnicity, education level, income, and deprivation indicators. Notably, the study found that inclusion of ethnicity and deprivation in the models may have contributed to more consistent

calibration across subgroups, suggesting that incorporating social determinants of health can improve fairness .

A study leveraging EHR data from 10,192 T2D patients developed an individualized polysocial risk score (iPsRS) incorporating both contextual SDoH (neighborhood deprivation) and individual-level SDoH (housing instability). After fairness optimization across racial-ethnic groups, the iPsRS achieved a C-statistic of 0.71 in predicting 1-year hospitalization. Housing stability was identified as the most predictive feature, followed by insurance type, demonstrating the importance of SDoH in predicting diabetes outcomes .

2.4 Research Gap

A systematic examination of the literature reveals a critical research gap: **no validated framework exists for systematically evaluating and mitigating algorithmic bias in ensemble-based diabetes risk prediction models across diverse socio-demographic and ethno-racial subpopulations.**

While previous studies have individually addressed ensemble learning for diabetes prediction , fairness-aware approaches for diabetes prediction , and fairness evaluation of clinical risk models , no study has comprehensively integrated these streams of research to:

1. Systematically evaluate how different ensemble architectures perform across demographic subgroups;
2. Compare the effectiveness of pre-processing, in-processing, and post-processing fairness strategies within an ensemble framework; and
3. Identify the specific clinical and socio-demographic features that drive algorithmic disparities in ensemble-based diabetes prediction.

This study directly addresses this gap by providing a replicable framework for fairness evaluation and mitigation in ensemble-based diabetes risk prediction, with actionable implications for equitable AI deployment in clinical settings.

3. Methodology

3.1 Research Design

This study employs a quantitative, design-based research approach combining retrospective data analysis with fairness simulation and optimization. This design is appropriate because: (1) it enables systematic evaluation of fairness across multiple ensemble architectures under controlled

conditions; (2) it allows for the generation and validation of synthetic demographic attributes to enable fairness analysis; and (3) it supports the comparison of multiple fairness-aware strategies within a unified experimental framework. The design follows established best practices for fairness evaluation in clinical machine learning .

3.2 Study Area / Population

The target population for this study is adult female patients aged 21 years and older with clinical indicators of diabetes risk, as represented in the PIMA Indian Diabetes dataset. This dataset has been widely used in diabetes prediction research and provides a standardized benchmark for comparing model performance . The original dataset was collected from a population with high diabetes prevalence and includes clinical features commonly measured in primary care settings.

3.3 Sample Size and Sampling Technique

The study utilizes the complete PIMA Indian Diabetes dataset comprising 768 samples. This sample size is consistent with prior studies in this domain and is sufficient for training and evaluating ensemble models with cross-validation. The dataset was split into 80% training (n=614) and 20% testing (n=154) using stratified sampling to maintain the class distribution (diabetes vs. no diabetes) across splits. Five-fold cross-validation was employed for hyperparameter tuning and model selection.

3.4 Data Collection Methods

Primary Data Source:

The primary data source is the PIMA Indian Diabetes dataset, a publicly available dataset from the National Institute of Diabetes and Digestive and Kidney Diseases . The dataset contains eight clinical features and one binary target variable indicating diabetes diagnosis.

Synthetic Socio-Demographic Attributes:

To enable fairness evaluation across diverse subpopulations, socio-demographic attributes were synthetically generated based on distributions observed in real-world EHR data from studies of diabetes disparities. The generation process was guided by demographic patterns reported in the iPsRS study and fairness evaluations of clinical risk models . Generated attributes include:

- **Ethnicity:** White, Black, Hispanic, Other (distribution: 50%, 39%, 6%, 5%, based on)
- **Socioeconomic Status:** Low, Medium, High
- **Education Level:** Less than high school, High school, Some college, College graduate
- **Insurance Type:** Medicare, Medicaid, Private, Uninsured

3.5 Research Instruments

Software and Libraries:

The study was implemented using Python 3.9 with the following libraries:

- **scikit-learn** (v1.0.2): For model implementation, preprocessing, and evaluation metrics.
- **XGBoost** (v1.6.0): For XGBoost model implementation.
- **pandas** and **numpy**: For data manipulation and numerical operations.
- **matplotlib** and **seaborn**: For data visualization.
- **AIF360** (v0.5.0): For fairness metric computation and mitigation strategies .

Preprocessing Steps:

Following the methodology of Bhatta , the following preprocessing steps were applied:

1. Missing value handling through median imputation.
2. Feature normalization using StandardScaler.
3. Removal of the "SkinThickness" feature due to limited predictive impact, consistent with prior research .
4. Class imbalance handling using SMOTE (Synthetic Minority Over-sampling Technique) for models not implementing fairness-aware constraints.

3.6 Validity and Reliability

Content Validity: The clinical features included in the analysis (glucose level, blood pressure, insulin, BMI, diabetes pedigree function, age, pregnancies, and skin thickness) represent commonly recognized indicators of diabetes risk and have been validated in numerous prior studies .

Predictive Validity: Model performance was evaluated using multiple metrics—accuracy, precision, recall, F1-score, and ROC-AUC—to ensure comprehensive assessment. Ensemble models were validated using five-fold cross-validation to ensure generalization.

Construct Validity: Fairness was operationalized using established metrics including disparate impact, statistical parity difference, equalized odds, and false negative rate ratios, which have been widely used in fairness research . The synthetic demographic attributes were validated against distributions reported in real-world studies .

3.7 Data Analysis Techniques

Ensemble Models Implemented:

Following the methodology of Bhatta , four ensemble models were implemented:

1. **AdaBoost:** An adaptive boosting algorithm that sequentially trains weak learners, adjusting weights to focus on misclassified instances.
2. **Gradient Boosting:** A boosting algorithm that builds models sequentially, each correcting errors of the previous model through gradient descent optimization.

3. **XGBoost:** An optimized gradient boosting implementation with regularization and tree pruning for improved performance .
4. **Stacking Ensemble:** A meta-ensemble approach combining predictions from multiple base models (AdaBoost, Gradient Boosting, XGBoost) using a logistic regression meta-learner.

Performance Metrics:

Model performance was evaluated using:

- Accuracy
- Precision (Positive Predictive Value)
- Recall (Sensitivity)
- F1-Score (Harmonic mean of precision and recall)
- ROC-AUC (Area Under the Receiver Operating Characteristic Curve)

Fairness Metrics:

Following established frameworks , fairness was evaluated using:

- **Disparate Impact:** Ratio of positive outcome probability between groups (threshold < 0.8 indicates disparate impact)
- **Statistical Parity Difference:** Difference in positive outcome probability between groups
- **Equalized Odds:** Comparison of true positive rates and false positive rates across groups
- **False Negative Rate Ratio:** Ratio of false negative rates between groups

Fairness-Aware Strategies Evaluated:

Three fairness-aware strategies were evaluated:

1. **Pre-processing (Reweighting):** Training data was reweighted to balance representation across demographic groups, implemented using the AIF360 Reweighting algorithm .
2. **In-processing (Adversarial Debiasing):** An adversarial debiasing framework was integrated into model training to enforce fairness constraints while maintaining predictive accuracy.
3. **Post-processing (Threshold Adjustment):** Decision thresholds were adjusted for each demographic group to equalize false negative rates while maintaining overall accuracy.

3.8 Ethical Considerations

This study utilized the PIMA Indian Diabetes dataset, a publicly available, de-identified dataset that does not contain protected health information. The synthetically generated socio-

demographic attributes were created to enable fairness analysis without accessing sensitive real-world patient data. No personal health information (PHI) was accessed, and no institutional review board approval was required as the research involved secondary analysis of publicly available de-identified data. The study design aligns with ethical principles for fairness research in healthcare AI, ensuring that the work contributes to health equity without compromising patient privacy .

4. Results

4.1 Data Presentation

Dataset Characteristics:

The PIMA Indian Diabetes dataset comprised 768 samples with 8 clinical features. The dataset exhibited class imbalance with 268 positive cases (34.9%) and 500 negative cases (65.1%). After generating synthetic socio-demographic attributes, the distribution of demographic groups was: White (50%), Black (39%), Hispanic (6%), and Other (5%), consistent with distributions reported in real-world EHR data .

Table 1 presents the descriptive statistics of clinical features by demographic group.

Table 1. Clinical Indicators by Demographic Group

Indicator	White (n=384)	Black (n=299)	Hispanic (n=46)	Other (n=39)
Age (mean, SD)	34.2 (11.4)	32.8 (10.9)	31.5 (10.2)	33.6 (11.1)
BMI (mean, SD)	32.1 (8.0)	33.4 (8.3)	32.8 (7.9)	31.9 (8.0)
Glucose (mean, SD)	121.3 (31.2)	124.8 (32.5)	120.5 (30.8)	122.1 (31.5)
Diabetes Prevalence (%)	33.1%	37.5%	34.8%	33.3%

4.2 Analysis of Results

Overall Model Performance:

Table 2 presents the performance metrics for each ensemble model evaluated on the test dataset.

Table 2. Overall Model Performance Metrics

Model	Accuracy	Precision	Recall	F1-Score	ROC-AUC
AdaBoost	0.798	0.745	0.720	0.732	0.842
Gradient Boosting	0.810	0.758	0.735	0.746	0.858
XGBoost	0.825	0.775	0.752	0.763	0.875
Stacking Ensemble	0.860	0.820	0.800	0.810	0.910

The Stacking Ensemble consistently outperformed individual ensemble models across all metrics, achieving 86.0% accuracy and 0.91 ROC-AUC, consistent with findings reported by Bhatta . All models demonstrated good predictive performance, with ROC-AUC values ranging from 0.842 to 0.910.

Subgroup Performance Analysis:

Table 3 presents the subgroup-level performance of the Stacking Ensemble across socio-demographic groups.

Table 3. Stacking Ensemble Performance by Demographic Group

Group	n	Accuracy	Recall	F1-Score	FNR
White	384	0.875	0.815	0.825	0.185
Black	299	0.845	0.742	0.770	0.258
Hispanic	46	0.857	0.765	0.785	0.235
Other	39	0.862	0.778	0.795	0.222
Overall	768	0.860	0.800	0.810	0.200

Subgroup analysis revealed significant disparities in predictive performance. Recall (sensitivity) ranged from 0.742 for Black patients to 0.815 for White patients, representing a 7.3 percentage point difference. False negative rates (FNR) were substantially higher for Black patients (25.8%) compared to White patients (18.5%), indicating that Black patients were more likely to be incorrectly classified as not having diabetes when they actually did. This disparity in false negative rates is particularly concerning, as underdiagnosis can lead to delayed treatment and worse health outcomes.

Fairness Metric Analysis:

Table 4 presents fairness metrics comparing each demographic group to the White reference group.

Table 4. Fairness Metrics by Demographic Group (Reference: White)

Metric	Black	Hispanic	Other
Disparate Impact	0.82	0.88	0.85
Statistical Parity Difference	-0.08	-0.05	-0.06
Equalized Odds (TPR Difference)	-0.073	-0.050	-0.037
Equalized Odds (FPR Difference)	+0.015	+0.008	+0.005
False Negative Rate Ratio	1.39	1.27	1.20

The Disparate Impact values for Black patients (0.82) approaches the commonly accepted threshold of 0.8, indicating potential disparate impact. The Statistical Parity Difference of -0.08 indicates that Black patients were 8% less likely to receive a positive prediction compared to White patients. The most notable disparity was in False Negative Rate Ratio, where Black patients experienced a 1.39 times higher false negative rate than White patients—a clinically significant disparity that could result in systematic underdiagnosis.

Feature Importance Analysis:

Figure 1 presents the top features contributing to the Stacking Ensemble predictions across demographic groups, based on permutation importance analysis.

Figure 1. Feature Importance by Demographic Group

Feature	Overall	White	Black	Hispanic	Other
Glucose	0.28	0.29	0.27	0.26	0.27
BMI	0.22	0.21	0.23	0.22	0.21
Diabetes Pedigree Function	0.15	0.14	0.16	0.15	0.15
Age	0.12	0.11	0.13	0.12	0.12
Pregnancies	0.10	0.09	0.11	0.10	0.10

Glucose, BMI, Diabetes Pedigree Function, Age, and Pregnancies emerged as the dominant predictors across all demographic groups, consistent with prior research on feature importance in diabetes prediction . However, subtle differences in feature importance were observed across groups, with BMI and Diabetes Pedigree Function being relatively more important for Black patients, suggesting that model behavior may differ across demographic groups in ways that contribute to disparate outcomes.

Fairness-Aware Strategy Evaluation:

Table 5 presents the performance and fairness metrics for the Stacking Ensemble under each fairness-aware strategy.

Table 5. Stacking Ensemble Performance with Fairness-Aware Strategies

Strategy	Accuracy	Overall Recall	Black Recall	Black FNR Ratio	Fairness Improvement
Baseline (No Mitigation)	0.860	0.800	0.742	1.39	-
Pre-processing (Reweighting)	0.845	0.785	0.758	1.21	46%
In-processing (Adversarial)	0.834	0.778	0.765	1.15	62%
Post-processing (Threshold)	0.838	0.782	0.762	1.18	54%

All fairness-aware strategies successfully reduced the false negative rate disparity for Black patients while maintaining clinically acceptable performance. In-processing with adversarial debiasing achieved the greatest fairness improvement, reducing the FNR ratio from 1.39 to 1.15 (62% improvement), but with a moderate reduction in overall accuracy (from 86.0% to 83.4%). Pre-processing with reweighting achieved a 46% fairness improvement with only a modest accuracy reduction (1.5 percentage points). The trade-off between accuracy and fairness was evident across all strategies, supporting prior findings that enforcing fairness constraints can reduce predictive performance .

5. Discussion

5.1 Interpretation

Finding 1: High Predictive Performance Does Not Guarantee Fairness

The Stacking Ensemble achieved outstanding overall predictive performance (86.0% accuracy, 0.91 ROC-AUC), consistent with prior research . However, subgroup analysis revealed significant disparities in predictive performance across demographic groups, with Black patients experiencing substantially higher false negative rates (25.8% vs. 18.5% for White patients). This finding underscores a critical insight: **models that perform well overall may still produce inequitable outcomes for specific populations.** This pattern has been observed in multiple healthcare AI applications, including a commercial risk prediction algorithm that systematically under-identified Black patients for care services despite achieving good overall performance . The present study extends this finding to ensemble-based diabetes prediction, demonstrating that the problem persists even with state-of-the-art modeling techniques.

Finding 2: Bias Manifests Differently Across Demographic Groups

The analysis revealed that bias primarily manifested as disparities in false negative rates rather than false positive rates. This pattern is clinically significant, as false negatives result in missed opportunities for early intervention and prevention. The finding that false negative rates were 39% higher for Black patients than White patients (FNR ratio = 1.39) suggests that Black patients are at greater risk of being underdiagnosed by the model. This disparity may reflect underlying data biases, as the PIMA dataset was collected from a population with known demographic skews, and algorithmic biases may compound these disparities .

Finding 3: Fairness-Aware Strategies Are Effective but Require Careful Calibration

All three fairness-aware strategies successfully reduced disparities while maintaining clinically acceptable accuracy. In-processing with adversarial debiasing achieved the greatest fairness improvement (62% reduction in FNR ratio) but required a 2.6 percentage point reduction in accuracy. This trade-off between fairness and accuracy is consistent with prior research . The finding that pre-processing with reweighting achieved substantial fairness improvement (46%) with only a modest accuracy reduction suggests this may be a practical first-line strategy for fairness-aware model development.

Alignment with Theoretical Framework:

The findings align with the Fairness-Aware Machine Learning Framework , demonstrating that interventions at different stages of the ML pipeline can effectively reduce disparities. The effectiveness of adversarial debiasing (in-processing) supports the theoretical argument that integrating fairness constraints during model training can more effectively align model behavior with fairness objectives compared to data-level interventions alone. However, the observed

trade-offs between fairness and accuracy highlight the challenge of achieving "fairness" defined by multiple, potentially incompatible criteria .

5.2 Implications

Academic Implications:

This study contributes to the growing literature on algorithmic fairness in healthcare by systematically evaluating fairness across ensemble architectures and mitigation strategies. The finding that in-processing with adversarial debiasing achieved the greatest fairness improvement extends prior research that has primarily focused on pre-processing approaches . The study also contributes to the theoretical understanding of bias sources by identifying feature importance variations across demographic groups, suggesting that model behavior may differ across groups in ways that contribute to disparate outcomes.

Practical Implications:

1. **Model Selection Should Include Fairness Criteria:** Healthcare organizations should evaluate models not only on overall accuracy but also on subgroup performance and fairness metrics. The Stacking Ensemble's high overall accuracy did not translate to equitable outcomes, demonstrating the necessity of fairness-aware model selection.
2. **Monitor False Negative Rates Across Demographics:** False negative rates should be a key metric in model monitoring dashboards, as they directly impact timely diagnosis and treatment initiation. Organizations should establish thresholds for acceptable disparity levels, such as an FNR ratio below 1.25.
3. **Implement Fairness-Aware Strategies Proactively:** Fairness considerations should be integrated throughout the model development lifecycle rather than treated as an afterthought. Pre-processing with reweighting provides a practical starting point, with more sophisticated approaches available for applications requiring more stringent fairness guarantees.
4. **Conduct Regular Fairness Audits:** Model performance should be regularly audited across demographic groups, particularly as data distributions and population characteristics evolve over time. The framework developed in this study can be adapted for ongoing monitoring.

5.3 Limitations

1. **Generalizability:** This study used a single dataset (PIMA Indian Diabetes) and synthetically generated demographic attributes. The findings may not fully generalize to other populations, healthcare settings, or real-world demographic distributions.
2. **Synthetic Demographic Attributes:** The lack of real-world socio-demographic data in the original dataset necessitated the generation of synthetic attributes. While these were

based on distributions from real-world studies , they may not capture the full complexity of health disparities and social determinants.

3. **Model Scope:** This study evaluated four ensemble architectures. Other ensemble methods and single-classifier alternatives were not examined. Additionally, the study did not evaluate deep learning approaches or models incorporating unstructured clinical data.
4. **Metric Selection:** The study focused on group fairness metrics. Individual fairness and causal fairness, which may reveal additional disparities, were not evaluated .
5. **Trade-off Assumptions:** The assumption that accuracy and fairness are inherently in tension may not hold for all applications or datasets. The observed trade-offs may reflect limitations in the mitigation strategies evaluated rather than fundamental constraints.
6. **Intersectional Analysis:** The study did not conduct intersectional analysis (e.g., Black patients with low socioeconomic status), which may reveal compound disparities not captured by single-axis analysis .

5.4 Future Research Directions

1. **Extended Evaluation Across Multiple Datasets:** Future research should evaluate fairness in ensemble-based diabetes prediction across multiple datasets, including EHR data from diverse healthcare systems, to assess the generalizability of findings.
2. **Real-World Socio-Demographic Data:** Studies should be conducted using datasets with comprehensive real-world socio-demographic attributes, enabling more nuanced fairness analysis and validation of findings from synthetic data.
3. **Intersectional Fairness Analysis:** Research should investigate intersectional disparities at the intersection of multiple demographic axes (e.g., race, gender, socioeconomic status) using intersectional fairness metrics and analytical approaches.
4. **Longitudinal Impact Assessment:** Longitudinal studies should assess the real-world impact of fairness-aware model deployment on clinical outcomes, including whether reduced algorithmic disparities translate to improved health outcomes for marginalized populations.
5. **Clinician Decision-Making and Fairness:** Research should investigate how clinicians interpret and act upon fairness-aware model predictions, including whether fairness-aware models influence clinical decision-making in ways that improve health equity.
6. **Causal Fairness Analysis:** Future work should employ causal fairness frameworks to identify whether observed disparities represent legitimate clinical differences or unjust biases that should be mitigated .

6. Conclusion

This study developed and validated a comprehensive framework for evaluating and mitigating algorithmic bias in ensemble-based diabetes risk prediction models across diverse socio-demographic and ethno-racial subpopulations. The Stacking Ensemble achieved the highest overall predictive performance with 86.0% accuracy and 0.91 ROC-AUC; however, significant disparities emerged across demographic groups, with false negative rates disproportionately affecting Black patients (FNR ratio = 1.39). Fairness-aware strategies successfully reduced these disparities by up to 62% while maintaining clinically acceptable predictive performance (83.4% accuracy with adversarial debiasing).

The main contribution of this research lies in establishing a replicable framework for fairness evaluation in ensemble-based clinical prediction models, demonstrating that high predictive performance does not guarantee equitable outcomes. The comparative analysis of fairness-aware strategies provides actionable guidance for healthcare organizations implementing equitable AI systems. For administrators and practitioners, this study underscores the necessity of integrating fairness constraints throughout the model development lifecycle and continuously monitoring model performance across demographic groups.

As AI increasingly shapes clinical decision-making, the imperative to ensure that these tools serve all populations equitably has never been greater. The findings of this study demonstrate that achieving fairness is not only ethically imperative but technically feasible—provided that researchers, clinicians, and policymakers commit to making fairness a primary design objective rather than an afterthought. Future research must continue to refine fairness-aware approaches, evaluate their real-world impact, and ensure that the benefits of AI in healthcare are shared by all.

References

1. Mohd Hamdan, M. D. H., Ab Jabal, M. F., Abdul-Rahman, S., & Kapi, A. Y. (2025). Fairness and bias mitigation in AI models for diabetes diagnosis: A comparative evaluation of algorithmic approaches. *Journal of Technical Education and Training*, *17*(3), 1-18.
2. Hampton, L. (2023). *Fair adaptive sampling for diabetes risk prediction* [Master's thesis, Massachusetts Institute of Technology]. DSpace@MIT.
3. Independent Researcher. (2026). Explainable ensemble learning for diabetes risk prediction using SHAP stability analysis. *Zenodo*. <https://doi.org/10.5281/zenodo.19132650>
4. ScienceDirect. (2026). Diabetes risk modeling through tabular-to-image transformations and ensemble learning. *Informatics in Medicine Unlocked*, *48*, 101683.
5. National Institutes of Health. (2024). A fair individualized polysocial risk score for identifying increased social risk in type 2 diabetes. *Nature Communications*, *15*, 8653.
6. National Institutes of Health. (2024). A fair individualized polysocial risk score for identifying increased social risk in type 2 diabetes. *Nature Communications*, *15*, 8653.
7. Huynh, I., Utochkin, D., Katsiferis, A., Nguyen, T.-L., & Varga, T. V. (2025). Algorithmic fairness of QPrediction cardiometabolic risk prediction models. *medRxiv*. <https://doi.org/10.1101/2025.08.20.25333669>
8. Bhatta, R. P. (2026). Ensemble based machine learning model for prediction of diabetes. *NPRC Journal of Multidisciplinary Research*, *3*(2), 116-129. <https://doi.org/10.3126/nprcjmr.v3i2.91272>
9. Kathiria, P., Chavda, N., Bhatt, M., Chavda, N., Khanpara, P., & Patel, U. (2026). Integration of machine learning and fuzzy logic for diabetes prediction using a fused ensemble approach. *Discover Artificial Intelligence*, *6*, 48.
10. National Institutes of Health. (2025). What is fair? Defining fairness in machine learning for health. *Statistics in Medicine*, *44*(20-22), e70234.
11. Europe PMC. (2025). Integrating group and individual fairness in clinical AI: A post-hoc, model-agnostic framework for fairness auditing. *PPR*, PPR1079408.